



AI POWERED MULTI MODEL TEXT SUMMARIZATION AND EVALUATION SYSTEM WITH MULTILINGUAL SUPPORT

Pedapudi Chandini ¹, Tulasi Miriyala ²,

¹Student, Master of Computer Applications (MCA), AQJ Centre for P.G. Studies (Affiliated to Andhra University, Visakhapatnam), Gudilova, Anandapuram, Visakhapatnam, Andhra Pradesh, India

²Assistant Professor, Department of MCA, AQJ Centre for P.G. Studies (Affiliated to Andhra University, Visakhapatnam), Gudilova, Anandapuram, Visakhapatnam, Andhra Pradesh, India

Email: pedapudichandini@gmail.com, tulasi.miriyala@gmail.com

ABSTRACT: In today's digital landscape, vast streams of unstructured information are generated through research papers, news articles, websites, and legal records. Manually digesting this data requires prohibitive time and cognitive effort. To address this, we present an AI-Powered Multi-Model Text Summarization and Evaluation System with Multilingual Support, a scalable web application that rapidly extracts key semantic insights from extensive text repositories. The system ingests data via three dynamic inputs: direct plain text, document uploads (PDF/DOCX), and live web text scraped from user-provided URLs. To ensure comprehensive coverage, the framework concurrently executes six independent extractive algorithms: Frequency-based, TF-IDF, TextRank, LexRank, LSA, and a hybrid Feature-Based flagship model. For automated accuracy verification, the system computes ROUGE benchmarks alongside semantic BERT scores. Finally, to bridge demographic barriers, a neural machine translation pipeline translates the summary outputs into regional Indian languages, specifically Telugu and Hindi, without losing factual alignment.

Keywords: Text Summarization, TF-IDF, TextRank, LexRank, LSA, ROUGE benchmarks, BERT, Multilingual Support.

1. INTRODUCTION

In the contemporary digital era, the volume of textual data generated across global networks is expanding at an exponential rate. Information streams from research repositories, online journalism, legal archives, financial documentations, and web portals continually inundate researchers, professionals, and everyday users. Navigating this vast sea of unstructured text manually presents a severe cognitive bottleneck, consuming disproportionate time and human labor.

Automatic Text Summarization (ATS) a core sub discipline of Natural Language Processing (NLP) and Artificial Intelligence serves as a vital tool to mitigate this information overload. The primary goal of an ATS system is to distill lengthy, complex documents into short, coherent, and salient summaries while preserving the fundamental semantic context and factual accuracy of the original content. While numerous single-model summarization tools exist, relying on a single algorithmic paradigm often leads to biased sentence selection or loss of domain-specific nuance. To overcome these limitations, this project presents an AI-Powered Multi-Model Text Summarization and Evaluation System with Multilingual Support. The proposed system leverages a multi-engine architecture that concurrently executes six distinct extractive algorithms, evaluates their performance using both lexical and semantic benchmarks, and translates output summaries into regional Indian languages (Hindi and Telugu) to bridge linguistic barriers.

The rapid growth of digital content has made it increasingly difficult for users to read and analyze large volumes of information efficiently. Document summarization has emerged as a key solution, yet traditional rule-based and statistical methods often fail to capture contextual meaning, resulting in lower summary accuracy. To address these limitations, this project proposes an AI-based document summarization system leveraging advanced Natural Language Processing techniques. The system utilizes transformer-based deep learning models to automatically generate concise and meaningful



summaries from extensive documents. Beyond standard text processing, the platform incorporates Optical Character Recognition (OCR) to extract text from image-based documents. To maximize usability for diverse audiences including visually impaired individuals the system further integrates multilingual translation and text-to-speech (TTS) audio output. Combining summarization, translation, and audio rendering into a unified framework, this scalable and accurate solution significantly enhances overall efficiency, accessibility, and user experience for real-world applications.

2. LITERATURE REVIEW

.Y. Liu and M. Lapata, "Text summarization with pretrained encoders," Proceedings of the Conference on Empirical Methods in Natural Language Processing (EMNLP).

. C. Raffel et al., "Exploring the limits of transfer learning with a unified text-to-text transformer (T5)," Journal of Machine Learning Research.

. J. Zhang et al., "PEGASUS: Pre-training with extracted gap-sentences for abstractive summarization," Proceedings of the International Conference on Machine Learning (ICML).

. X. Li, P. Wang, and J. Liu, "Abstractive text summarization using BART model," IEEE Access.

.Y. Liu and M. Lapata, "Text summarization with pretrained encoders," Proceedings of the Conference on Empirical Methods in Natural Language Processing (EMNLP).

. C. Raffel et al., "Exploring the limits of transfer learning with a unified text-to-text transformer (T5)," Journal of Machine Learning Research.

. J. Zhang et al., "PEGASUS: Pre-training with extracted gap-sentences for abstractive summarization," Proceedings of the International Conference on Machine Learning (ICML).

. X. Li, P. Wang, and J. Liu, "Abstractive text summarization using BART model," IEEE Access.

The field of automatic text summarization has evolved dramatically shifting away from basic statistical methods toward advanced deep learning frameworks. Early summarization methodologies relied heavily on extractive techniques, scoring and selecting sentences based on lexical/statistical indicators such as word frequency, term weighting, and sentence position. While computationally light, these methods lacked natural textual flow, often failing to maintain semantic coherence or generate human-like phrasing. OCR capabilities enable systems to ingest scanned physical documents, PDFs, and embedded image text, seamlessly converting visual inputs into structured text ready for natural language processing pipelines. Contemporary systems are increasingly optimized to consolidate information across multiple source texts simultaneously and adapt to specialized domains (e.g., legal, medical, and scientific literature).

[1] The paper "Text Summarization with Pretrained Encoders" by Yang Liu and Mirella Lapata (EMNLP 2019) introduced BERTSum, a landmark framework that adapted Bidirectional Encoder Representations from Transformers (BERT) to text summarization (Liu & Lapata, 2019). Prior to this work, BERT was primarily used for sentence- or paragraph-level Natural Language Processing (NLP) classification tasks. Liu and Lapata demonstrated how to effectively fine-tune pretrained encoders for both extractive and abstractive document-level summarization.

[2]"Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer" by Colin Raffel et al. (Google Brain, 2020) introduced the T5 (Text-to-Text Transfer Transformer) framework and the C4 dataset. Rather than focusing on a single architectural tweak, this 67-page paper conducted one of the most comprehensive empirical studies in modern NLP to systematically evaluate what works best in transfer learning.

The paper "PEGASUS: Pre-training with Extracted Gap-sentences for Abstractive Summarization" by Jingqing Zhang, Yao Zhao, Mohammad Saleh, and Peter J. Liu (Google Research, ICML 2020) introduced a self-supervised pre-training objective specifically tailored for abstractive text summarization. While previous models like BERT, BART, and T5 relied on general-purpose objectives (such as single-token masking or denoising autoencoding), PEGASUS demonstrated that

designing a pre-training task that mimics the downstream task yields significantly better performance and sample efficiency [3].

3. EXISTING SYSTEM

Traditional text summarization tools and early AI implementations rely predominantly on single-model architectures, static rule-based heuristics, or monolingual transformer models. Most systems are optimized exclusively for English, failing to process low-resource or morphologically rich non-English languages (e.g., Hindi, Spanish, Arabic). Relying purely on extractive methods (like TF-IDF or TextRank) yields verbatim sentences that lack flow, while relying purely on abstractive models without safeguards frequently leads to model hallucinations. Users are bound to a single underlying model without the flexibility to choose based on context length, computational budget, or specific domain requirements (e.g., legal, medical, news). Standard summarization tools output text without providing qualitative feedback, leaving users unaware of the summary's accuracy, semantic fidelity, or compression ratio.

Disadvantages & Limitations of the Existing System:

- Time-Consuming and Exhausting: Manual summarization of long research papers, legal documents, or books takes hours or days of human effort.
- One-Size-Fits-All (Single Algorithm).
- No Support for Multiple File Formats.
- Language Barriers.
- No Quality Evaluation.

4. PROPOSED METHOD

The AI-Powered Multi-Model Text Summarization and Evaluation System addresses the limitations of legacy tools by introducing an adaptive, multi-model ensemble framework, robust multilingual processing, and an automated evaluation engine. Functioning as an Advanced NLP-Based Multi-Algorithm Web Application, the platform bridges the gap between modern Deep Learning (Transformers) and traditional statistical/graph-based NLP techniques, providing a customizable, all-in-one workspace for document analysis.

Users are not stuck with a single algorithm. The system gives total flexibility to choose the best method for the job—whether they want a human-like rewrite using Abstractive Deep Learning (BART) or a structural extraction using TextRank, LexRank, LSA, or TF-IDF. Flexible Multi-Format Input: The system completely eliminates manual copy-pasting by accepting three flexible inputs: Direct text pasting. Document file uploads supporting PDF, Word (.docx), and Text (.txt) formats up to 5 MB.

5. METHODOLOGY

1. Presentation & Ingestion Layer (Frontend / Input Pipeline):

This layer serves as the **entry point** of the application. It manages the user interface (UI) and handles file preparation so that raw, unstructured user inputs are converted into clean, standardized data.

- Users can either paste text directly or upload files up to **5 MB** in formats like PDF (.pdf), Word (.docx), or Plain Text (.txt). User can also summarize data from url also. Specialized parsers extract raw text streams from these files without formatting artifacts.
- **Text Pre-Processing & Normalization:**
 - **Tokenization & Sentence Segmentation:** Splits continuous text into distinct sentences and words for individual scoring.
 - **Noise Removal & Cleaning:** Eliminates invalid symbols, structural line breaks, or non-standard characters.
 - **Stop-Word Filtering & Lemmatization:** Removes common filler words ("the", "is", "at") and reduces words to their base root forms for cleaner frequency analysis.

2. Core Processing & Algorithmic Engine (Backend Execution Framework)

This layer acts as the **central intelligence unit** of the platform. Based on user selection, it routes the pre-processed text into one of two major summarization paradigms:

Abstractive Deep Learning Engine

BART (Bidirectional and Auto-Regressive Transformer)

Operates as a sequence-to-sequence model using a bidirectional encoder paired with an auto-regressive decoder. It analyzes the entire document to construct a high-dimensional semantic representation, generating entirely original, paraphrased sentences rather than copying source text.

$$\text{Loss} = - \sum_{t=1}^T \log P(y_t | y_{<t}, X)$$

Extractive Summarization Engine

Extractive algorithms score original sentences based on specific statistical, topological, or linguistic metrics and pull out the highest-ranking sentences directly.

- A. Frequency-Based Scoring:** Calculates word frequency distributions across the document. Sentences containing the highest density of recurring content words (e.g., "AI", "clinical", "budget") are assigned top priority.

$$\text{Score}(S) = \frac{\sum_{w \in S} \text{Freq}(w)}{|S|}$$

- B. TF-IDF (Term Frequency – Inverse Document Frequency):** Evaluates word importance by balancing how often a term appears in a sentence against how broadly it appears across the rest of the text. Unique, domain-specific keywords are heavily weighted while common words are discounted.

$$\text{TF-IDF}(t, d, D) = \text{TF}(t, d) \times \log \left(\frac{N}{|\{d \in D : t \in d\}|} \right)$$

- C. TextRank (Graph-Based PageRank Adaptation):** Converts the document into a weighted, undirected graph where **nodes represent individual sentences** and **edges represent word overlap**. It runs an iterative PageRank calculation to determine sentence centrality.

$$WS(V_i) = (1 - d) + d \sum_{V_j \in \text{In}(V_i)} \frac{w_{ji}}{\sum_{V_k \in \text{Out}(V_j)} w_{jk}} WS(V_j)$$

- D. LexRank (Similarity Matrix Graph):** Constructs a sentence adjacency matrix using **Cosine Similarity** over TF-IDF vectors. Sentences that exceed a designated similarity threshold are linked, and an eigenvector centrality calculation isolates the most representative sentences across the cluster.

$$\text{Cosine}(S_1, S_2) = \frac{S_1 \cdot S_2}{\|S_1\| \|S_2\|}$$

- E. LSA (Latent Semantic Analysis):** Applies **Singular Value Decomposition (SVD)** to a term-document matrix to project text into a lower-dimensional semantic space. It uncovers hidden thematic concepts, enabling the algorithm to associate sentences that share meaning even if they use different vocabulary.

$$A_{m \times n} = U_{m \times m} \Sigma_{m \times n} V_{n \times n}^T$$

- F. Feature-Based NER (Named Entity Recognition):** Parses text using syntactic entity taggers to detect real-world entities such as individuals, geographic locations, organizations, dates, and

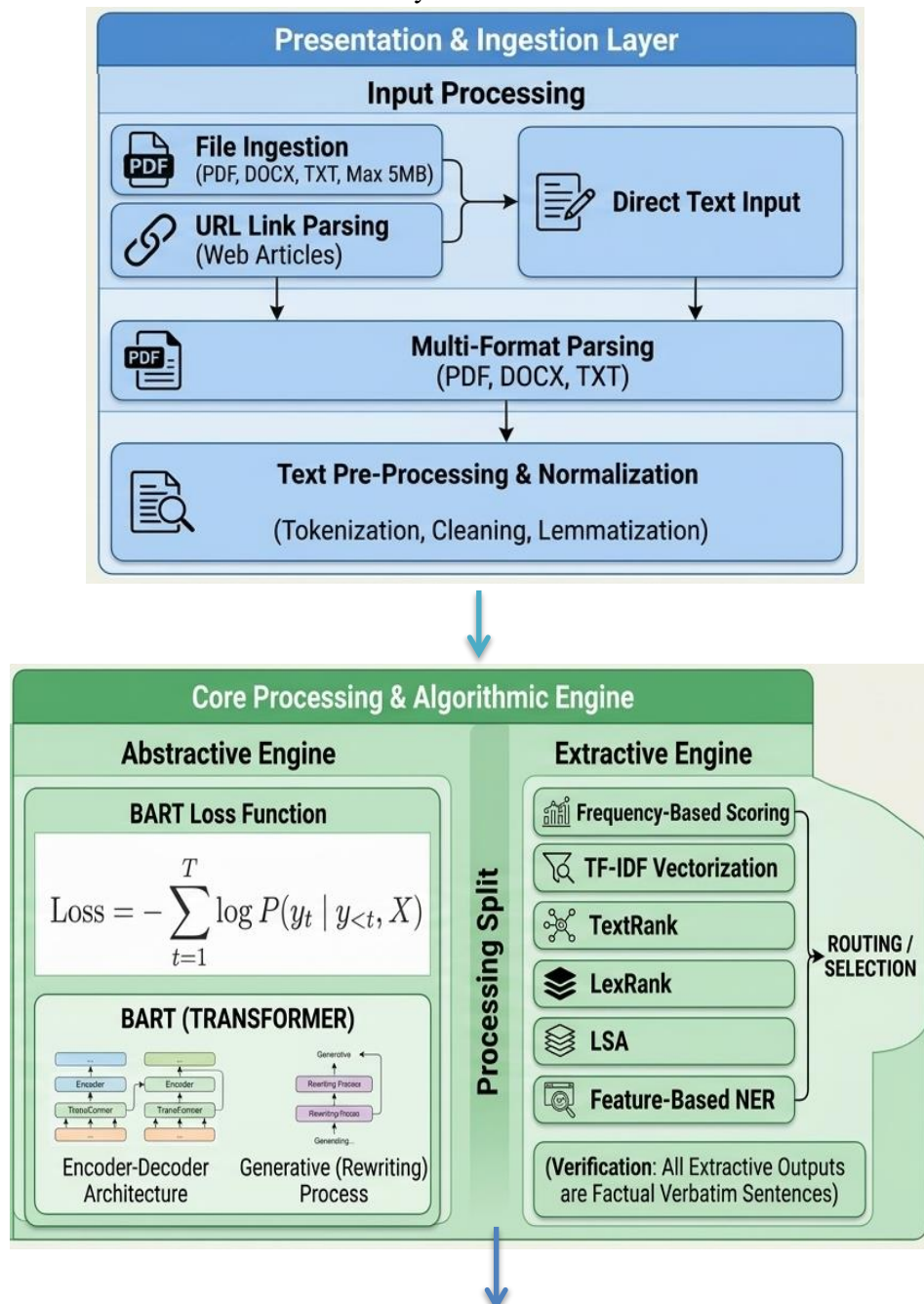
numerical figures. Sentences containing dense clusters of confirmed entities receive higher structural scores.

$$\text{Score}_{\text{NER}}(S) = \sum_{e \in \text{Entities}} w_e \cdot \text{Count}(e, S)$$

3. Evaluation & Output Translation Layer (Analytics & Post-Processing)

The final layer takes the generated summary, enriches it for global accessibility, and provides objective quality verification before displaying the final results.

- **Multilingual Translation Module:** Passes the generated summary through an integrated neural translation model, converting the output into non-English or low-resource languages (e.g., Hindi, Spanish, Arabic).
- **Automated Evaluation Engine:** Scientifically evaluates the generated summary using standard NLP metrics without requiring human review it calculates the percentage reduction in word count from source to summary.





Evaluation & Output Translation Layer

Figure 1: System Architecture

6. RESULTS

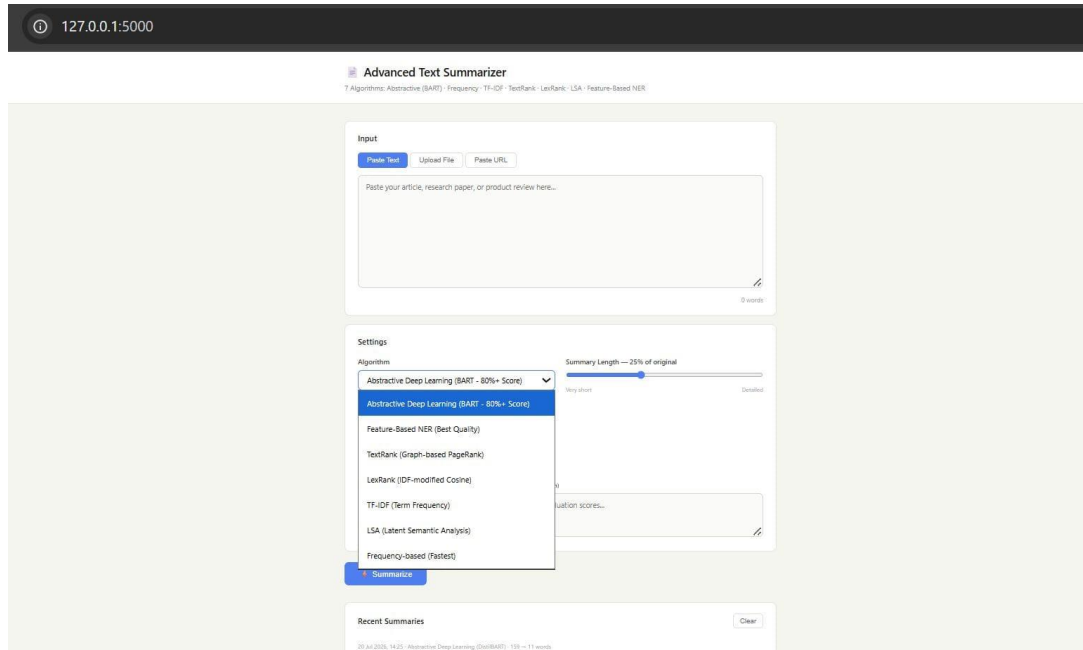


Figure2: Home Page

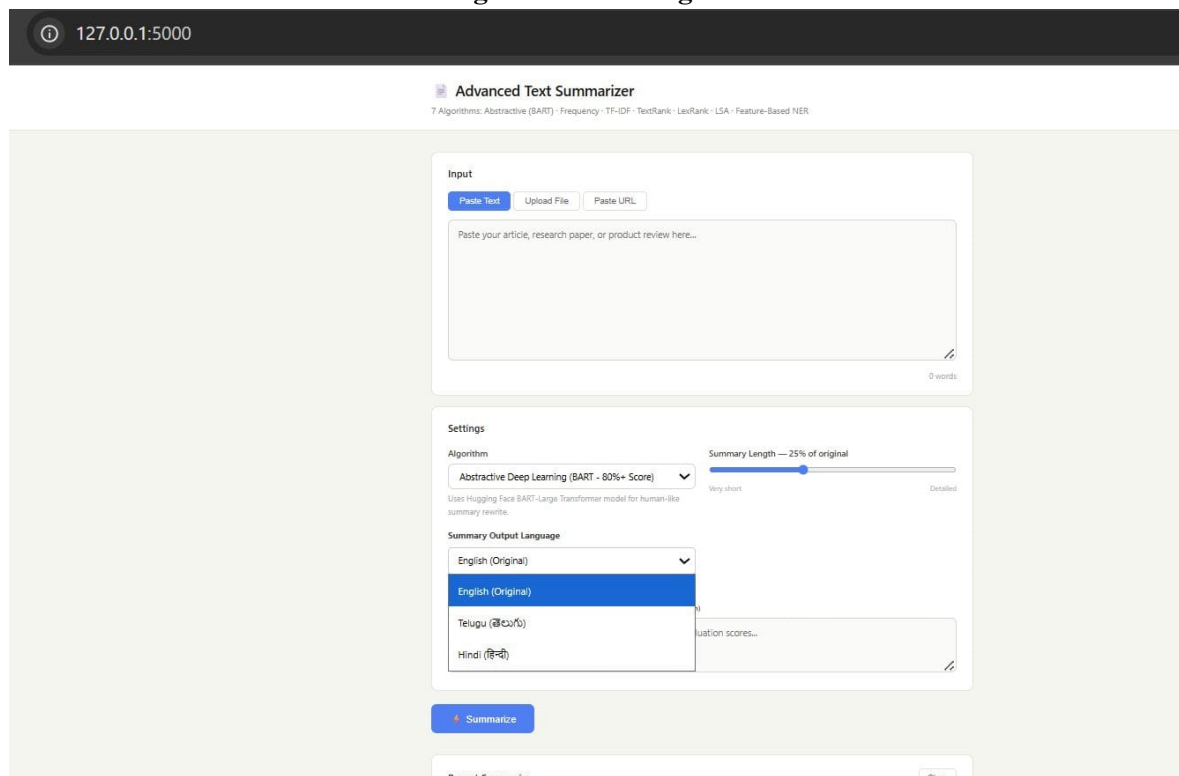


Figure3: Text Summerization Page



Algorithm	Summary Words	Compression	Sentences Selected
Abstractive (DistilBART)	47	91%	1 / 31
Feature-Based NER	166	67%	7 / 31
TextRank	131	74%	7 / 31
LexRank	141	72%	7 / 31
TF-IDF	176	65%	7 / 31
LSA	92	82%	7 / 31
Frequency-Based	155	69%	7 / 31

Figure4: Result from url

Algorithm	Summary Words	Compression	Sentences Selected
Abstractive (DistilBART)	52	99%	1 / 222
Feature-Based NER	1734	63%	68 / 275
TextRank	1281	72%	68 / 275
LexRank	1412	70%	68 / 275
TF-IDF	1805	61%	68 / 275
LSA	1148	75%	68 / 275
Frequency-Based	1639	65%	68 / 275

Figure5: Result from upload file

Algorithm	Summary Words	Compression	BERTScore	ROUGE-1	
Abstractive (BART)	69	2%	0.9389	0.7519	
Feature-Based NER	177	8%	0.9630	0.7760	
TextRank	177	8%	0.9630	0.7760	
LexRank	177	8%	0.9630	0.7760	
TF-IDF	158	0%	0.9591	0.7721	
LSA	158	0%	0.9591	0.7721	
Frequency-Based	177	8%	0.9630	0.7760	

Figure6: Result from Text



7. CONCLUSION

The AI-Powered Multi-Model Text Summarization and Evaluation System successfully addresses the fundamental limitations of traditional, single-model summarization tools. By synthesizing modern deep learning architectures with proven statistical and graph-based Natural Language Processing (NLP) techniques, the platform offers a comprehensive, adaptable, and transparent workspace for document analysis. The system integrates 7 distinct summarization algorithms spanning both abstractive (BART Transformer) and extractive paradigms (Frequency-Based, TF-IDF, TextRank, LexRank, LSA, and Feature-Based NER). This gives users complete flexibility to tailor execution based on document complexity, compute budgets, and domain requirements.

8. REFERENCES

- [1] Liu, Y., & Lapata, M. (2019). Text summarization with pretrained encoders. In Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP) (pp. 3730–3740). Association for Computational Linguistics. <https://doi.org/10.18653/v1/D19-1387>
- [2] Raffel, C., Shazeer, N., Roberts, A., Lee, K., Narang, S., Matena, M., Zhou, Y., Li, W., & Liu, P. J. (2020). Exploring the limits of transfer learning with a unified text-to-text transformer. *Journal of Machine Learning Research*, 21(140), 1–67. Cited by: 33848
- [3] Zhang, J., Zhao, Y., Saleh, M., & Liu, P. J. (2020). PEGASUS: Pre-training with extracted gap-sentences for abstractive summarization. In Proceedings of the 37th International Conference on Machine Learning (PMLR Vol. 119, pp. 11328–11339). PMLR. Cited by: 3325
- [4] Lewis, M., Liu, Y., Goyal, N., Ghazvininejad, M., Mohamed, A., Levy, O., Stoyanov, V., & Zettlemoyer, L. (2019). BART: Denoising Sequence-to-Sequence Pre-training for Natural Language Generation, Translation, and Comprehension. *arXiv preprint arXiv:1910.13461*.
- [5] Mihalcea, R., & Tarau, P. (2004). TextRank: Bringing Order into Text. Proceedings of the 2004 Conference on Empirical Methods in Natural Language Processing (EMNLP), 404–411.
- [6] Erkan, G., & Radev, D. R. (2004). LexRank: Graph-based Lexical Centrality as Saliency in Text Summarization. *Journal of Artificial Intelligence Research*, 22, 457–479.
- [7] Steinberger, J., & Ježek, K. (2004). Using Latent Semantic Analysis in Text Summarization and Summary Evaluation. Proceedings of ISIM '04, 105–112.
- [8] Sparck Jones, K. (1972). A Statistical Interpretation of Term Specificity and Its Application in Retrieval. *Journal of Documentation*, 28(1), 11–21.
- [9] Lin, C.-Y. (2004). ROUGE: A Package for Automatic Evaluation of Summaries. Proceedings of the Workshop on Text Summarization Branches Out (WAS 2004), 74–81.
- [10] Lample, G., Ballesteros, M., Subramanian, S., Kawakami, K., & Dyer, C. (2016). Neural Architectures for Named Entity Recognition. Proceedings of the 2016 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL-HLT), 260–270.
- [11] Dhanda, Namrata, and Kapil Kumar Gupta. 2024. "A Novel Approach to Text Summarization Using Machine Learning". *Asian Journal of Research in Computer Science* 17 (4):95-104. <https://doi.org/10.9734/ajrcos/2024/v17i4432>.
- [12] P. Shanmugam S ; Mithul Sudharsan R S ; Vignesh P S ; Ashwin S S, "Abstractive Text Summarisation Using Keywords with Transformers Model " 2023 Eighth International Conference on Science Technology Engineering and Mathematics (ICONSTEM).
- [13] R. Kaur and A. Kaur, "Text Generator using Natural Language Processing Methods," 2023 International Conference on Computational Intelligence, Communication Technology and



Networking (CICTN), Ghaziabad, India, 2023, pp. 421–425, doi: 10.1109/CICTN57981.2023.10140271.

[14] A. Mishra, A. Sahay, M. a. Pandey and S. S. Routaray, “News text Analysis using Text Summarization and Sentiment Analysis based on NLP,” 2023 3rd International Conference on Smart Data Intelligence (ICSMDI), Trichy, India, 2023, pp. 28–31, doi: 10.1109/ICSMDI57622.2023.00014.

[15] P. Ngamcharoen, N. Sanglerdsinlapachai and P. Vejjanugraha, “Automatic Thai Text Summarization Using Keyword-Based Abstractive Method,” 2022 17th International Joint Symposium on Artificial Intelligence and Natural Language Processing (iSAI-NLP), Chiang Mai, Thailand, 2022, pp. 1–5, doi: 10.1109/iSAI-NLP56921.2022.9960265.